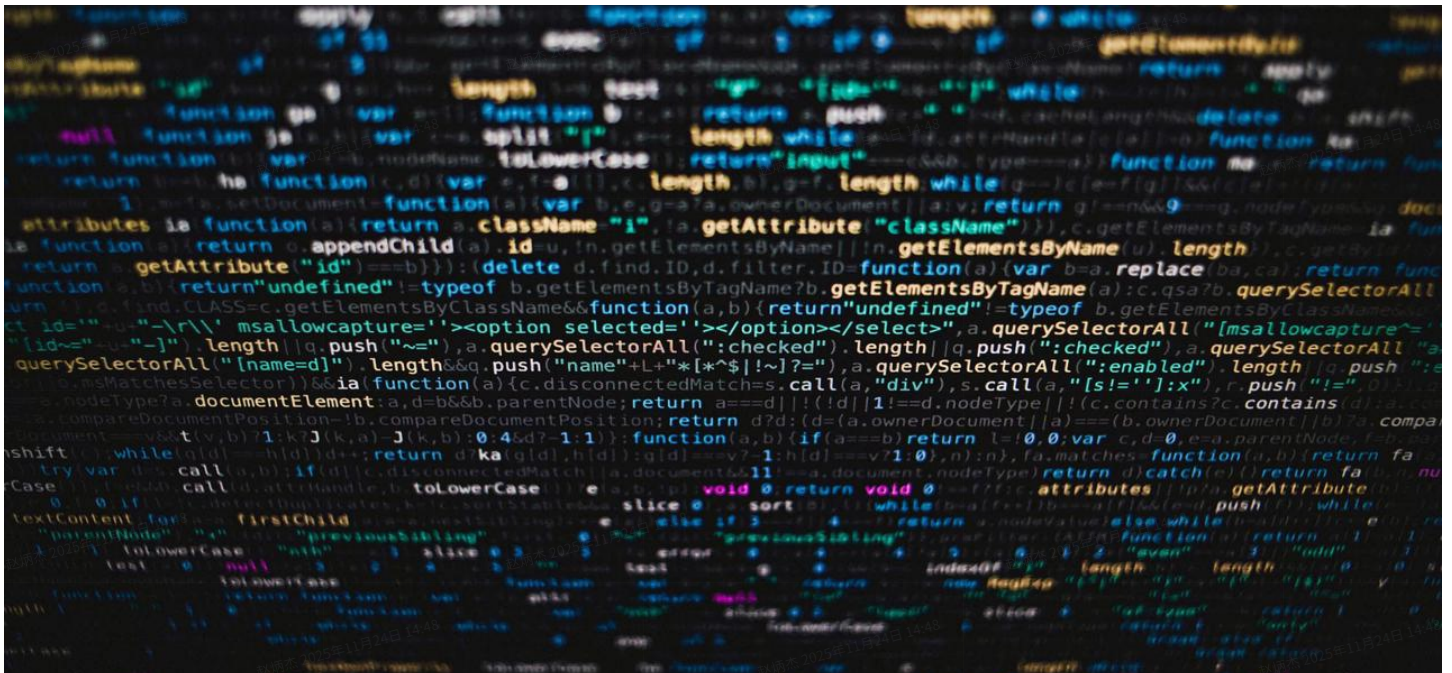




Swift框架微调

1. 环境准备

推荐安装方式

1. 创建conda环境

```
bash
1  conda create -n ljn_swift python=3.11 -y
2  conda activate ljn_swift
```

2. 安装CUDA12.8

```
bash
1  sudo apt-get -y install cuda-toolkit-12-8
```

3. 安装对应的torch版本

```
bash
1  pip3 install torch torchvision --index-url
https://download.pytorch.org/whl/cu128
```

4. 安装ms-swift依赖

```
bash
1 pip install 'ms-swift[all]' -U
```

这些版本目前（2025.11.14）可以运行。如果不行，可能是ms-swift更新所致。

如果以上步骤不行，建议参考这里一步步安装：[Swift DOCUMENTATION — swift 3.10.0.dev0 文档](#)

当前环境

服务器：hsmy-aigc-algorithm-gpu-ecs-dev-01

服务器ip：45.78.200.93

93机器上的可用环境名：ljn_swift

我的导出：



full_environment.yml
4.92KB



*此文件仅供参考，最好使用前面的步骤安装。

2. 数据集准备

自定义数据集格式要求

[自定义数据集 — swift 3.10.0.dev0 文档](#)

文档中提到，swift支持多种格式数据集。我们的数据集是视频-QA对，使用jsonl格式。

基本格式如下：

代码块

```
1 {"messages": [{"role": "user", "content": "浙江的省会在哪? "}, {"role":  
  "assistant", "content": "浙江的省会在杭州。"}]}
```

```
2 {"messages": [{"role": "user", "content": "<image><image>两张图片有什么区别"},  
  {"role": "assistant", "content": "前一张是小猫，后一张是小狗"}], "images":  
  ["/xxx/x.jpg", "/xxx/x.png"]}
```

```
3 {"messages": [{"role": "user", "content": "<audio>语音说了什么"}, {"role":  
  "assistant", "content": "今天天气真好呀"}], "audios": ["/xxx/x.mp3"]}
```

```
4 {"messages": [{"role": "system", "content": "你是有个有用无害的助手"}, {"role":  
  "user", "content": "<image>图片中是什么, <video>视频中是什么"}, {"role":  
  "assistant", "content": "图片中是一个大象，视频中是一只小狗在草地上奔跑"}], "images":  
  ["/xxx/x.jpg"], "videos": ["/xxx/x.mp4"]}
```

meta-qa数据集

由于解析原因，部分视频不完整或者损坏。因此先对数据集进行清洗，将这部分视频及其QA对删除。

清洗代码路径：

清洗视频文件： `/data/aigc/ljn/data_preprocess/meta-filter.py`

清洗数据jsonl格式QA对： `/data/aigc/ljn/data_preprocess/meta-dataset-filter.py`

数据路径：

视频文件： `/data/aigc/deliancen/workpace/AIGC/Data/violation`

QA

对： `/data/aigc/deliancen/workpace/AIGC/Data/violation/mp4_annotation`

数据集： `/data/aigc/ljn/data_preprocess/data-to-train/meta_qa_pairs_cleaned.jsonl`

meta-qa数据集文件：



meta_qa_pairs_cleaned.jsonl

62.15MB



测试用数据集

测试使用的数据集是MVbench的action_sequence数据，就是对动作的描述。原来的数据是选择题，我做了转换，存在/data/aigc/ljn/swift/action_sequence_converted.jsonl了。

以下是截取的部分数据集：

`/data/aigc/ljn/swift/action_sequence_converted.jsonl`

- ```
{ "messages": [{ "role": "user", "content": "<video> What happened after the person took the food? Choices: Ate the medicine.; Tidied up the blanket.; Put down the cup/glass/bottle.; Took the box." }, { "role": "assistant", "content": "Ate the medicine." }], "videos": ["/data/aigc/bj/Qwen3_VL_Measurement/datasets/MVBench/video/star/star/Charades_v1_480/ZS9XR.mp4"] }
```
- ```
{ "messages": [ { "role": "user", "content": "<video> What happened before the person watched at the book? Choices: Tidied up the table.; Took the phone/camera.; Opened the closet/cabinet.; Washed the table." }, { "role": "assistant", "content": "Opened the closet/cabinet." } ], "videos": [ "/data/aigc/bj/Qwen3_VL_Measurement/datasets/MVBench/video/star/star/Charades_v1_480/EY6P4.mp4" ] }
```

更严谨一点的话，应该把choices也去掉。

3. 运行训练

训练前准备

记得先 `chmod +x train_swift_2.sh`，给训练脚本赋权。

路径：`/data/aigc/ljn/swift/train_swift_2.sh`

训练脚本 (meta-qa)

train_swift_2.sh

```
1  PYTORCH_CUDA_ALLOC_CONF='expandable_segments:True' \  
2  IMAGE_MAX_TOKEN_NUM=300 \  
3  VIDEO_MAX_TOKEN_NUM=300 \  
4  NPROC_PER_NODE=8 \  
5  FPS=1 \  
6  FPS_MAX_FRAMES=192 \  
7  CUDA_VISIBLE_DEVICES=0,1,2,3,4,5,6,7 \  
8  swift sft \  
9      --model /data/aigc/ljn/models_to_test/Qwen3-VL-8B-Instruct \  
10     --dataset '/data/aigc/ljn/data_preprocess/data-to-  
    train/meta_qa_pairs_cleaned.jsonl' \  
11     --load_from_cache_file False \  
12     --split_dataset_ratio 0.01 \  
13     --train_type full \  
14     --torch_dtype bfloat16 \  
15     --num_train_epochs 2 \  
16     --per_device_train_batch_size 1 \  
17     --per_device_eval_batch_size 1 \  
18     --learning_rate 1e-5 \  
19     --freeze_vit true \  
20     --freeze_aligner true \  
21     --gradient_checkpointing true \  
22     --vit_gradient_checkpointing true \  
23     --gradient_accumulation_steps 4 \  
24     --eval_steps 100 \  
25     --save_steps 100 \  
26     --save_total_limit 2 \  
27     --logging_steps 5 \  
28     --output_dir /data/aigc/ljn/swift/tec-chi-swifted/meta-qa \  
29     --warmup_ratio 0.05 \  
30     --deepspeed zero3 \  
31     --dataset_num_proc 8 \  

```

```
32 --report_to swanlab \  
33 --swanlab_project meta-qa \  
34 --dataloader_num_workers 4 \  
35 --resume_from_checkpoint /data/aigc/ljn/swift/tec-chi-swifted/meta-qa/v15-  
20251107-181750/checkpoint-200
```

技巧与注意事项

1. 模型在较小的数据集上如果设置了较高的学习率，那么很容易导致过拟合。具体来说：对于200条左右的10s以内的视频，学习率设置为 $1e-4$ ，训练5个epoch时会出现模型的回答**全部是一样的** token的情况，这就代表过拟合：因为模型会取概率最高的token作为下一个输出，梯度收敛于这个token了，自然就会全是一样的。

之后考虑初始学习率设置为 $1e-5$ 。因为存在

2. 因为 `deepspeed` 的gpu管理和 `torch` 的 `device_map` 功能会有冲突，导致必须在训练脚本中显式指定 `NPROC_PER_NODE = $你的gpu数量`，不然会报如下错误：

```
bash
```

```
1 ValueError: DeepSpeed is not compatible with device_map. n_gpu: 4,  
  local_world_size: 1.
```

3. `--max_length` 这里指定为8192的话，会因为视频过长而导致tokenize之后超过这个长度，报错：

```
bash
```

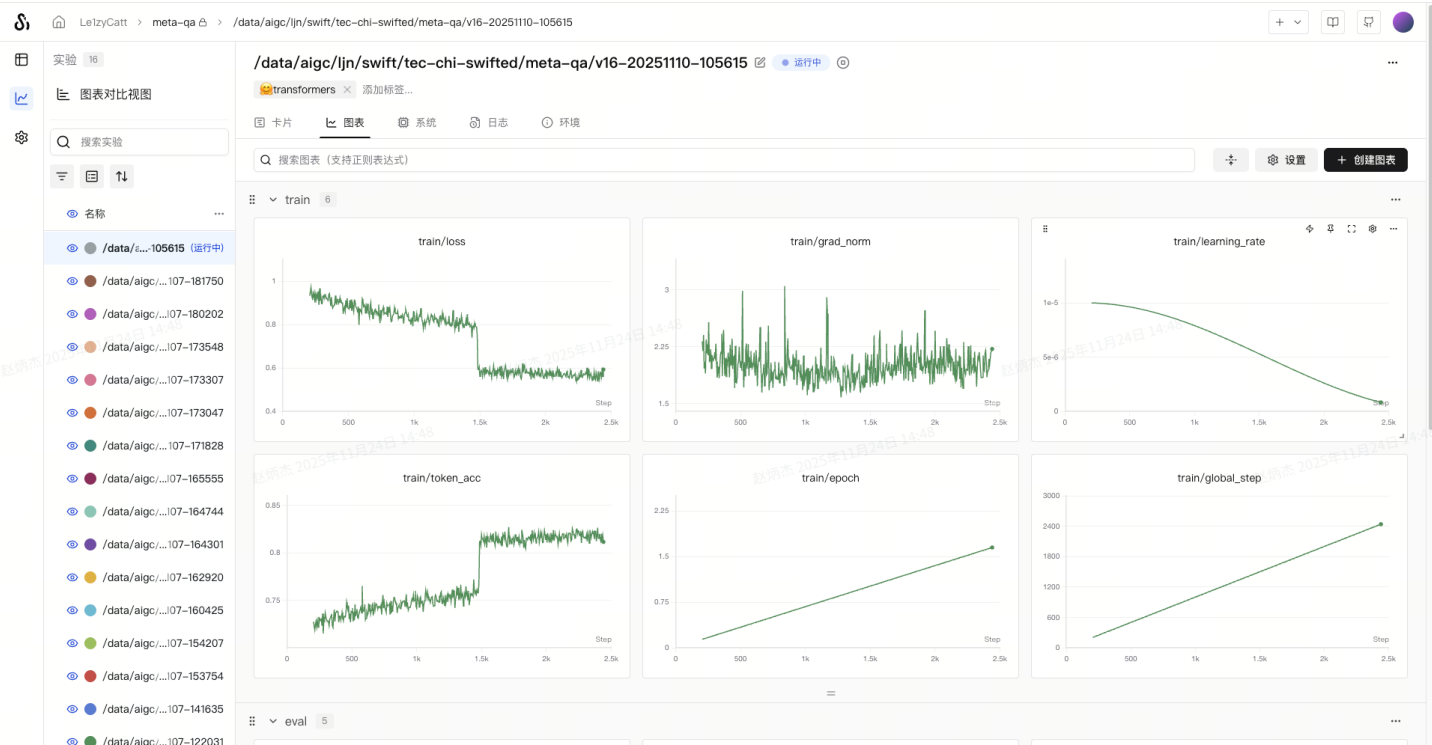
```
1 swift.llm.template.base.MaxLengthError: Current length of row(10744) is larger  
  than the max_length(8192).  
2  
3 [WARNING:swift] 🙌🙌🙌 There are errors in the template.encode, and another  
  piece of data will be randomly selected.
```

虽然报错之后还是会继续训练，但用的是随机的别的数据。所以要写大一点，具体写多少应该在有了数据集之后，抽样做tokenize，然后观察一下分布再做决定。我在 `train_swift_2.sh` (meta-qa训练脚本) 里删除了这一参数，根据文档：

🔥 `max_length`: 限制单数据集样本经过 `tokenizer.encode` 后的tokens最大长度，超过的数据样本会根据 `truncation_strategy` 参数进行处理（避免训练OOM）。默认为None，即设置为模型支持的tokens最大长度(`max_model_len`)。

这里的最大token数应该是Qwen3-VL-8B-Instruct支持的最长上下文长度，即**256K tokens**。

4. 加入参数 `--report_to swanlab` 可以使用swanlab界面，需配合 `--swanlab_project` 参数指定项目名称。第一次使用需自行登录swanlab获取密钥。查看meta-qa的训练记录找我  罗景楠。



5. 如上图，loss在中间出现突降。原因是训练设置 `--num_train_epochs 2`，总共2958steps，所以会在 $2958/2 = 1479$ step进入第二个epoch，因此会出现loss突降，属于正常情况。

4. 参数说明

5. 训练历史

meta违规素材qa对训练

脚本改动

- i. 输入长度上限 `--max_length 16384` 参数删除。使用默认值，即模型最大输入（256K token）。
- ii. 最大输入帧数 `FPS_MAX_FRAMES=192`，每秒帧数 `FPS=1`。

训练结果

